

# Research - Small Model Efficiency

Alpasan, Lois-Kleinner

*Lois-Kleinner and 0-1.gg 2026 - Inte11ect Platform Documentation Confidential  
- All Rights Reserved*

---

#  
re-  
search  
-  
Doc-  
u-  
ment  
01

---

Small  
Model  
Effi-  
ciency

---

>  
**As-**  
**so-**  
**ci-**  
**ated**  
**Mod-**  
**ule:**  
Tec-  
11 >  
**Cat-**  
**e-**  
**gory:**  
Re-  
search  
&  
De-  
vel-  
op-  
ment  
>  
**Last**  
**Up-**  
**dated:**  
2026-  
06-  
19  
##  
Ab-  
stract

---

This  
doc-  
u-  
ment  
presents  
a  
com-  
pre-  
hen-  
sive  
anal-  
ysis  
of  
small  
model  
effi-  
ciency  
within  
the  
Intellect  
plat-  
form  
ar-  
chi-  
tec-  
ture,  
fo-  
cus-  
ing  
on  
the  
de-  
ploy-  
ment  
of  
Qwen2-  
VL-  
2B  
as  
the  
pri-  
mary  
vision-  
language  
back-  
bone.  
The  
in<sup>3</sup>  
ves-  
tiga-  
tion  
cov-  
ers  
quan-  
tiza-

---

##  
1.  
In-  
tro-  
duc-  
tion

---

The rapid advancement of large language models has created an increasing demand for efficient inference solutions that can operate within constrained computational environments. The Intellect platform<sup>5</sup> addresses this challenge through a

---

Small  
model  
effi-  
ciency  
is  
not  
merely  
a  
mat-  
ter  
of  
pa-  
ram-  
eter  
count  
re-  
duc-  
tion;  
it  
en-  
com-  
passes  
a  
holis-  
tic  
opti-  
miza-  
tion  
of  
the  
en-  
tire  
in-  
fer-  
ence  
stack,  
in-  
clud-  
ing  
model  
ar-  
chi-  
tec-  
ture,  
quan-  
tiza-  
tion  
pre-  
ci-  
sion,  
mem-  
ory  
man-  
age-  
ment,

---

The  
cur-  
rent  
land-  
scape  
of  
effi-  
cient  
AI  
in-  
fer-  
ence  
is  
char-  
ac-  
ter-  
ized  
by  
sev-  
eral  
com-  
pet-  
ing  
paradigms.  
Model  
com-  
pres-  
sion  
through  
quan-  
tiza-  
tion  
has  
emerged  
as a  
dom-  
i-  
nant  
strat-  
egy,  
with  
tech-  
niques  
rang-  
ing  
from  
post-  
training  
quan-  
tiza-  
tion  
(PTQ)  
to  
quantization-  
aware

---

This document is organized as follows: Section 2 reviews the architectural foundations of the Qwen2-VL-2B integration. Section 3 presents the quantization framework and empirical results. Section 4 analyzes inference



---

##  
2.  
Ar-  
chi-  
tec-  
tural  
Foun-  
da-  
tions  
###  
2.1  
Qwen2-  
VL-  
2B  
Ar-  
chi-  
tec-  
ture  
Overview

---

The Qwen2-VL-2B model, developed by Alibaba Cloud's Qwen team, represents a significant advancement in efficient vision-language modeling. The architecture employs a modified transformer decoder with the following specifications:

---

```
python
#
Qwen2-VL-2B
configuration
structure
model_config
= {
    "hidden_size":
    1536,
    "intermediate_size":
    8960,
    "num_attention_heads":
    12,
    "num_hidden_layers":
    24,
    "num_key_value_heads":
    4,
    "vocab_size":
    151936,
    "rope_theta":
    1000000,
    "max_position_embeddings":
    32768,
    "vision_hidden_size":
    1024,
    "vision_num_channels":
    3,
    "vision_patch_size":
    14,
    "vision_image_size":
    448
}
```

---

The  
model  
pro-  
cesses  
vi-  
sual  
in-  
puts  
through  
a  
Vi-  
sion  
Trans-  
former  
(ViT)  
back-  
bone  
with  
 $14 \times 14$   
patch  
em-  
bed-  
ding,  
pro-  
ject-  
ing  
vi-  
sual  
fea-  
tures  
into  
the  
lan-  
guage  
model's  
hid-  
den  
space  
via  
a  
cross-  
modal  
pro-  
jec-  
tion  
layer.  
This  
de<sup>12</sup>  
sign  
en-  
ables  
effi-  
cient  
multi-  
modal

---

###  
2.2  
Intellect  
Mod-  
ule  
Inte-  
gra-  
tion

---

Within  
the  
Intellect  
plat-  
form,  
Qwen2-  
VL-  
2B  
is  
dis-  
tributed  
across  
the  
72-  
module  
ar-  
chi-  
tec-  
ture  
through  
the  
Eigen-  
vec-  
tor  
Rout-  
ing  
mech-  
a-  
nism  
(GOD-  
11).  
Each  
mod-  
ule  
spe-  
cial-  
izes  
in a  
sub-  
set  
of  
the  
model's  
ca-  
pa-  
bili-  
ties:

—  
|  
Mod-  
ule  
Group

|  
Mod-  
ule  
IDs

|  
Spe-  
cial-  
iza-  
tion

| Pa-  
ram-  
eter  
Al-  
loca-  
tion

| |—

|—

|—

|—|

| Vi-  
sual  
Cor-  
tex

|  
Tec-  
11,  
Read-  
11,  
Asc-  
11 |

Im-  
age  
en-  
cod-  
ing,  
fea-  
ture  
ex-  
trac-  
tion

|  
480M

| |  
15  
Lan-  
guage  
Pro-  
cess-  
ing |

Taut-  
11,  
Muse-

---

###  
2.3  
Struc-  
tured  
Spar-  
sity  
The  
plat-  
form  
im-  
ple-  
ments  
struc-  
tured  
spar-  
sity  
pat-  
terns  
that  
ex-  
ploit  
the  
mod-  
ular  
ar-  
chi-  
tec-  
ture  
to  
re-  
duce  
com-  
pu-  
ta-  
tional  
re-  
quire-  
ments:



---

```

“python
class
Struc-
turedSpar-
sity-
Router:
def
init(self,
mod-
ule__count=72,
spar-
sity__ratio=0.6):
self.module__count
=
mod-
ule__count
self.sparsity__ratio
=
spar-
sity__ratio
self.routing_matrix
=
torch.zeros((module__count,
mod-
ule__count))

```

---

```

def
compute_routing_path(self,
input_embedding):
#
Eigenvector-
based
routing
with
sparsity
mask
relevance_scores
=
self._compute_relevance(input_embedding)
top_k
=
int(self.module_count
* (1 -
self.sparsity_ratio))
mask
=
torch.topk(relevance_scores,
top_k).indices
return
mask
““

```

---

The  
struc-  
tured  
spar-  
sity  
ap-  
proach  
yields  
a  
60%  
re-  
duc-  
tion  
in  
acti-  
vated  
pa-  
ram-  
e-  
ters  
per  
in-  
fer-  
ence  
step,  
di-  
rectly  
con-  
tribut-  
ing  
to  
the  
plat-  
form's  
effi-  
ciency  
gains.

---

###  
2.4  
Hi-  
erar-  
chi-  
cal  
At-  
ten-  
tion  
Mech-  
a-  
nism

---

Beyond  
stan-  
dard  
at-  
ten-  
tion,  
the  
Intellect  
plat-  
form  
im-  
ple-  
ments  
a hi-  
erar-  
chi-  
cal  
at-  
ten-  
tion  
mech-  
a-  
nism  
that  
par-  
ti-  
tions  
the  
at-  
ten-  
tion  
com-  
pu-  
ta-  
tion  
across  
mul-  
ti-  
ple  
reso-  
lu-  
tion  
lev-  
els:

---

```

“python
class
Hi-
erar-
chi-
ca-
lAt-
ten-
tion(torch.nn.Module):
def
init(self,
hid-
den_size,
num_heads,
num_levels=3):
su-
per().__init__()
self.hidden_size
=
hid-
den_size
self.num_heads
=
num_heads
self.num_levels
=
num_levels
self.head_dim
=
hid-
den_size
//
num_heads
self.level_projections
=
torch.nn.ModuleList([
torch.nn.Linear(hidden_size,
hid-
den_size
//
(2 **
i))
for i
in
range(num_levels)])

```

---

```

def
for-
ward(self,
hid-
den_states,
at-
ten-
tion_mask=None):
out-
puts
= []
for
level
in
range(self.num_levels):
pro-
jected
=
self.level_projections[level]
attn_output
=
flash_attn.flash_attn_func(
pro-
jected,
pro-
jected,
pro-
jected,
dropout_p=0.0,
soft-
max_scale=(self.head_dim
//
(2 **
level))
**)
-0.5,
causal=True
)
out-
puts.append(attn_output)
re-
turn
torch.cat(outputs,
dim=-
1)
““

```

---

This  
hier-  
ar-  
chi-  
cal  
ap-  
proach  
re-  
duces  
the  
asymptotic  
complexity  
of  
attention  
from  
 $O(n^2)$   
to  
 $O(n \log n)$   
for  
typical  
sequence  
lengths,  
providing  
substantial  
speedups  
for  
long-context  
scenarios.



---

###  
2.5  
Cross-  
Modal  
Pro-  
jec-  
tion  
De-  
sign  
The  
cross-  
modal  
pro-  
jec-  
tion  
layer  
map-  
ping  
vi-  
sual  
fea-  
tures  
to  
the  
lan-  
guage  
space  
em-  
ploys  
a  
multi-  
layer  
per-  
cep-  
tron  
with  
gated  
acti-  
va-  
tion:

---

```

“python
class
Cross-
Modal-
Pro-
jec-
tion(torch.nn.Module):
def
init(self,
vi-
sion_dim=1024,
lan-
guage_dim=1536,
hid-
den_dim=2048):
su-
per().__init__()
self.vision_proj
=
torch.nn.Linear(vision_dim,
hid-
den_dim)
self.language_proj
=
torch.nn.Linear(language_dim,
hid-
den_dim)
self.gate
=
torch.nn.Linear(hidden_dim,
hid-
den_dim)
self.output
=
torch.nn.Linear(hidden_dim,
lan-
guage_dim)

```

---

```

def
for-
ward(self,
vi-
sion_features,
lan-
guage_features):
v =
torch.nn.functional.silu(self.vision_proj(vision_features))
l =
self.language_proj(language_features)
g =
torch.sigmoid(self.gate(v))
fused
= g
* v
+ (1
- g)
* l
re-
turn
self.output(fused)
““

```

---

This  
gated  
fu-  
sion  
mech-  
a-  
nism  
al-  
lows  
the  
model  
to  
dy-  
nam-  
i-  
cally  
bal-  
ance  
vi-  
sual  
and  
lin-  
guis-  
tic  
in-  
for-  
ma-  
tion  
based  
on  
the  
in-  
put  
con-  
text,  
im-  
prov-  
ing  
cross-  
modal  
rea-  
son-  
ing  
qual-  
ity  
while  
main-  
tain-  
ing  
com-  
pu-  
ta-  
tional  
effi-  
ciency.

---

##  
3.  
Quan-  
tiza-  
tion  
Frame-  
work  
###  
3.1  
Pre-  
ci-  
sion  
Tier  
Strat-  
egy

---

The  
Intellect  
plat-  
form  
em-  
ploys  
a  
tiered  
quan-  
tiza-  
tion  
strat-  
egy  
that  
adapts  
pre-  
ci-  
sion  
based  
on  
mod-  
ule  
sen-  
si-  
tiv-  
ity  
and  
task  
re-  
quire-  
ments:

—  
|  
Tier  
|  
Pre-  
ci-  
sion  
|  
Bit  
Width  
|  
Mod-  
ules  
| Ac-  
cu-  
racy  
Im-  
pact  
| |—  
|—  
|—  
|—  
|—|  
|  
Crit-  
ical  
|  
FP16  
| 16  
bits  
|  
GOD-  
11,  
Arch-  
11 |  
<0.1%  
| |  
Stan-  
dard  
|  
INT8  
| 8  
bits  
|  
Tec-  
11,  
Read-  
11,  
31  
Acc-  
11 |  
<0.5%  
| |  
Effi-  
cient  
|  
INT4

---

###  
3.2  
Post-  
Training  
Quan-  
tiza-  
tion  
Im-  
ple-  
men-  
ta-  
tion  
“python  
im-  
port  
torch  
from  
torch.ao.quantization  
im-  
port  
quan-  
tize\_dynamic,  
Quant-  
Type



---

```

def
ap-
ply_quantization_pipeline(model,
pre-
ci-
sion_map):
quan-
tized_modules
= {}
for
mod-
ule_name,
mod-
ule
in
model.named_modules():
if
mod-
ule_name
in
pre-
ci-
sion_map:
pre-
ci-
sion
=
pre-
ci-
sion_map[module_name]
if
pre-
ci-
sion
==
"INT8":
quan-
tized_modules[module_name]
=
quan-
tize_dynamic(
mod-
ule,
{torch.nn.Linear},
dtype=torch.qint8
)
elif
pre-
ci-
sion
==
"INT4":
quan-
tized_modules[module_name]
=

```

---

###  
3.3  
Cal-  
ibra-  
tion  
and  
Vali-  
da-  
tion

---

The  
quan-  
tiza-  
tion  
cali-  
bra-  
tion  
pro-  
cess  
uti-  
lizes  
1,024  
rep-  
re-  
sen-  
ta-  
tive  
sam-  
ples  
from  
the  
eval-  
ua-  
tion  
dataset  
to  
de-  
ter-  
mine  
opti-  
mal  
scal-  
ing  
fac-  
tors.  
The  
vali-  
da-  
tion  
pro-  
to-  
col  
mea-  
sures  
both  
nu-  
mer-  
ical  
di-  
ver-  
gence  
and  
task-  
specific  
per-

```

mermaid
graph TD
    A[Input Data] --> B[Calibration Dataset]
    B --> C[Quantization Calibration]
    C --> D{Validation}
    D -->|Pass| E[Deploy Quantized Model]
    D -->|Fail| F[Adjust Precision Tier]
    F --> C
    E --> G[Production Inference]
    G --> H[Monitor Drift]
    H -->|Drift| C
    H -->|Threshold| C

```

---

###  
3.4  
Per-  
Channel  
vs  
Per-  
Tensor  
Quan-  
tiza-  
tion

---

A  
crit-  
ical  
de-  
sign  
deci-  
sion  
in  
the  
quan-  
tiza-  
tion  
frame-  
work  
is  
the  
gran-  
ular-  
ity  
of  
quan-  
tiza-  
tion  
scal-  
ing  
fac-  
tors.  
Per-  
channel  
quan-  
tiza-  
tion  
as-  
signs  
sep-  
a-  
rate  
scal-  
ing  
fac-  
tors  
to  
each  
out-  
put  
chan-  
nel  
of a  
lin-  
ear  
layer,  
while  
per-  
tensor  
quan-

---

```

“python
def
com-
pare_quantization_strategies(weight_matrix):
#
Per-
tensor
quan-
tiza-
tion
scale_max
=
weight_matrix.abs().max()
per_tensor_quant
=
(weight_matrix
/
scale_max
*
127).round().to(torch.int8)
#
Per-
channel
quan-
tiza-
tion
chan-
nel_max
=
weight_matrix.abs().max(dim=1,
keep-
dim=True).values
per_channel_quant
=
(weight_matrix
/
chan-
nel_max
*
127).round().to(torch.int8)

```

---

```

quantization_error
= {
    "per_tensor_mse":
    torch.mean((weight_matrix
    -
    per_tensor_quant.float()
    *
    scale_max
    /
    127)
    **
    2).item(),
    "per_channel_mse":
    torch.mean((weight_matrix
    -
    per_channel_quant.float()
    *
    channel_max
    /
    127)
    **
    2).item()
} re-
turn
quan-
tiza-
tion_error
““

```



---

Experimental

re-

sults

indi-

cate

that

per-

channel

quan-

tiza-

tion

re-

duces

the

mean

squared

er-

ror

by

ap-

prox-

i-

mately

$3.5\times$

com-

pared

to

per-

tensor

ap-

proaches,

with

min-

imal

ad-

di-

tional

com-

pu-

ta-

tional

over-

head

dur-

ing

in-

fer-

ence.

---

###  
3.5  
Sen-  
si-  
tiv-  
ity  
Anal-  
ysis  
for  
Mixed-  
Precision  
Al-  
loca-  
tion  
The  
allo-  
ca-  
tion  
of  
pre-  
ci-  
sion  
tiers  
across  
mod-  
ules  
is  
guided  
by a  
sys-  
tem-  
atic  
sen-  
si-  
tiv-  
ity  
anal-  
ysis:

---

mermaid  
flowchart  
TD  
A[Sensitivity  
Analysis  
Pipeline]  
-->  
B[Module  
Isolation]  
B  
-->  
C[Quantize  
Each  
Module  
Independently]  
C  
-->  
D[Measure  
Perplexity  
Impact]  
D  
-->  
E[Measure  
Task  
Accuracy  
Impact]  
E  
-->  
F[Compute  
Sensitivity  
Score]  
F  
-->  
G{Rank  
Modules}  
G  
-->  
H[High  
Sensitivity:  
FP16]  
G  
-->  
I[Medium  
Sensitivity:  
INT8]  
G  
-->  
J<sup>4</sup><sub>4</sub>Low  
Sensitivity:  
INT4/NF4]  
H  
-->  
K[Deploy  
Mixed-Precision  
Model]

---

The  
 sen-  
 si-  
 tiv-  
 ity  
 score  
 for  
 each  
 mod-  
 ule  
 is  
 com-  
 puted  
 as:

$$\text{Sensitivity}(m) = \frac{|\Delta \text{PPL}_m|}{|\text{PPL}_{\text{baseline}}| + |\Delta \text{Acc}_m|}$$

---

Where  $\Delta\text{PPL}_m$  is the perplexity increase when module  $m$  is quantized,  $\Delta\text{Acc}_m$  is the accuracy decrease, and  $\alpha_m$  is a weighting factor typically set to 2.0 to emphasize task performance preservation.

---

###  
3.6  
Quantization-  
Aware  
Fine-  
Tuning  
For  
mod-  
ules  
where  
post-  
training  
quan-  
tiza-  
tion  
re-  
sults  
in  
un-  
ac-  
cept-  
able  
ac-  
cu-  
racy  
degra-  
da-  
tion,  
the  
plat-  
form  
sup-  
ports  
quantization-  
aware  
fine-  
tuning  
(QAFT):

---

```

“python
class
Quan-
tiza-
tion-
Aware-
Fine-
Tuner:
def
init(self,
model,
quan-
tized_modules,
learning_rate=5e-
5):
self.model
=
model
self.quantized_modules
=
quan-
tized_modules
self.optimizer
=
torch.optim.AdamW(
[p
for
m in
quan-
tized_modules
for
p in
m.parameters()],
lr=learning_rate
)

```

---

```

def
train_step(self,
batch):
#
For-
ward
pass
with
straight-
through
esti-
ma-
tor
for
mod-
ule
in
self.quantized_modules:
mod-
ule.apply(self._fake_quantization)
outputs
=
self.model(**batch)
loss
=
out-
puts.loss
#
Back-
ward
pass
(gra-
di-
ents
by-
pass
quan-
tiza-
tion)
loss.backward()
self.optimizer.step()
self.optimizer.zero_grad()
return
loss.item()

```



---

```

def
_fake_quantization(self,
module):
if
instance(module,
torch.nn.Linear):
module.weight.data
=
self._quantize_dequantize(module.weight.data)
““

```

---

QAFT  
re-  
cov-  
ers  
an  
av-  
er-  
age  
of  
1.3%  
ac-  
cu-  
racy  
across  
the  
seven  
eval-  
ua-  
tion  
bench-  
marks,  
with  
the  
most  
sig-  
nifi-  
cant  
im-  
prove-  
ments  
ob-  
served  
in  
GSM8K  
(+2.1%)  
and  
Hu-  
manEval  
(+1.8%).

---

##  
4.  
In-  
fer-  
ence  
Pipeline  
Op-  
ti-  
miza-  
tion  
###  
4.1  
Batched  
Pro-  
cess-  
ing  
Ar-  
chi-  
tec-  
ture  
The  
in-  
fer-  
ence  
pipeline  
im-  
ple-  
ments  
dy-  
namic  
batch-  
ing  
with  
priority-  
aware  
schedul-  
ing:

---

```
“‘rust
//
Tauri-
based
in-
fer-
ence
en-
gine
schedul-
ing
pub
struct
In-
fer-
enceSched-
uler
{
batch__size:
usize,
max__latency__ms:
u64,
pri-
or-
ity__queues:
HashMap<Priority,
VecD-
eque>,
}
```

---

```

impl
In-
fer-
enceSched-
uler
{
pub
fn
build_batch(&mut
self)
->
Vec
{ let
mut
batch
=
Vec::with_capacity(self.batch_size);
let
mut
cur-
rent_size
= 0;

```

---

```

//
Priority-
based
batch
con-
struc-
tion
for
pri-
or-
ity
in
[Pri-
or-
ity::High,
Pri-
or-
ity::Normal,
Pri-
or-
ity::Low]
{
if let
Some(queue)
=
self.priority_queues.get_mut(&priority)
{
while
cur-
rent_size
<
self.batch_size
{
if let
Some(request)
=
queue.pop_front()
{
cur-
rent_size
+=
re-
quest.token_estimate;
batch.push(request);
}
else
{
break;
} } }
}
batch
} }
““

```

---

###  
4.2  
KV-  
Cache  
Op-  
ti-  
miza-  
tion  
Key-  
value  
cache  
man-  
age-  
ment  
is  
crit-  
ical  
for  
effi-  
cient  
au-  
tore-  
gres-  
sive  
gen-  
era-  
tion.  
The  
Intellect  
plat-  
form  
im-  
ple-  
ments  
a  
slid-  
ing  
win-  
dow  
cache  
with  
importance-  
based  
evic-  
tion:

---

```

“python
class
Slid-
ing-
Win-
dowKV-
Cache:
def
init(self,
win-
dow_size=4096,
max_cache_size=8192):
self.window_size
=
win-
dow_size
self.max_cache_size
=
max_cache_size
self.cache
= {}
self.importance_scores
=
{}
def
up-
date(self,
layer_id,
key,
value,
po-
si-
tion_ids):
if
layer_id
not
in
self.cache:
self.cache[layer_id]
=
{“keys”:
[],
“val-
ues”:
[]}
```



---

```
self.cache[layer_id]["keys"].append(key)
self.cache[layer_id]["values"].append(value)
#
Evic-
tion
when
ex-
ceed-
ing
max
cache
size
if
len(self.cache[layer_id]["keys"])
>
self.max_cache_size:
self._evict(layer_id)
```

---

```

def
    _evict(self,
    layer_id):
    #
    Score-
    based
    eviction
    of
    least
    im-
    por-
    tant
    to-
    kens
    scores
    =
    self.importance_scores.get(layer_id,
    []) if
    scores:
    evict_indices
    =
    torch.topk(
    torch.tensor(scores),
    k=len(scores)
    -
    self.window_size,
    largest=False
    ).in-
    dices
    for
    idx
    in
    sorted(evict_indices,
    re-
    verse=True):
    del
    self.cache[layer_id]["keys"][idx]
    del
    self.cache[layer_id]["values"][idx]
    ““

```

---

###  
4.3  
Flash  
At-  
ten-  
tion  
Inte-  
gra-  
tion  
The  
plat-  
form  
lever-  
ages  
Flash  
At-  
ten-  
tion  
v2  
to  
re-  
duce  
mem-  
ory  
com-  
plex-  
ity  
from  
 $O(n^2)$   
to  
 $O(n \log n)$   
for  
at-  
ten-  
tion  
com-  
pu-  
ta-  
tions:  
“python  
im-  
port  
flash\_attn

---

```

class
FlashAt-
ten-
tion-
Layer(torch.nn.Module):
def
init(self,
hid-
den_size,
num_heads):
su-
per().__init__()
self.hidden_size
=
hid-
den_size
self.num_heads
=
num_heads
self.head_dim
=
hid-
den_size
//
num_heads

```

---

```

def
for-
ward(self,
query,
key,
value,
mask=None):
#
Flash
at-
ten-
tion
with
memory-
efficient
tiling
out-
put
=
flash_attn.flash_attn_func(
query,
key,
value,
dropout_p=0.0,
soft-
max_scale=self.head_dim
**)
-0.5,
causal=True
) re-
turn
out-
put
“‘

###
4.4
Con-
tin-
u-
ous
Batch-
ing
Strat-
egy

```

---

Beyond  
static  
dy-  
namic  
batch-  
ing,  
the  
pipeline  
im-  
ple-  
ments  
con-  
tin-  
u-  
ous  
batch-  
ing  
where  
new  
re-  
quests  
are  
ad-  
mit-  
ted  
as  
ear-  
lier  
ones  
com-  
plete:

---

```

“rust
pub
struct
Con-
tin-
u-
ous-
Batcher
{ ac-
tive_requests:
Vec,
pend-
ing_queue:
VecD-
eque,
max_concurrent:
usize,
sched-
uler:
Arc<Mutex>,
}
impl
Con-
tin-
u-
ous-
Batcher
{
pub
fn
run_iteration(&mut
self)
->
Vec
{ let
mut
com-
pleted
=
Vec::new();

```

```

//
Check
for
com-
pleted
re-
quests
self.active_requests.retain(|req|
{ if
req.is_complete()
{
com-
pleted.push(req.into_completed());
false
}
else
{
true
} }));
//
Fill
slots
with
pend-
ing
re-
quests
let
slots
=
self.max_concurrent
-
self.active_requests.len();
for
_ in
0..slots
{
if let
Some(pending)
=
self.pending_queue.pop_front()
{
self.active_requests.push(ActiveRequest::new(pending));
} }

```



completed  
} }  
““

Continuous  
batch-  
ing  
in-  
creases  
over-  
all  
through-  
put  
by  
15-  
25%  
com-  
pared  
to  
tra-  
di-  
tional  
static  
batch-  
ing,  
par-  
ticu-  
larly  
un-  
der  
vari-  
able  
load  
con-  
di-  
tions.

---

###  
4.5  
Op-  
era-  
tor  
Fu-  
sion  
and  
Ker-  
nel  
Op-  
ti-  
miza-  
tion  
The  
in-  
fer-  
ence  
en-  
gine  
em-  
ploys  
op-  
era-  
tor  
fu-  
sion  
to  
re-  
duce  
ker-  
nel  
launch  
over-  
head  
and  
im-  
prove  
mem-  
ory  
lo-  
cal-  
ity:

---

```

python
def
fuse_operations(computation_graph):
fused_ops
=
[]
i =
0
while
i <
len(computation_graph):
current
=
computation_graph[i]
#
Try
to
fuse
consecutive
element-wise
operations
if
current.type
==
"element_wise":
fused
=
current
j =
i +
1
while
j <
len(computation_graph)
and
computation_graph[j].type
==
"element_wise":
fused
=
fuse_element_wise(fused,
computation_graph[j])
j
+=
1
fused_ops.append(fused)
i =
j67
else:
fused_ops.append(current)
i
+=
1
return
fused_ops

```

---

Operator  
fu-  
sion  
re-  
duces  
the  
num-  
ber  
of  
ker-  
nel  
launches  
by  
ap-  
prox-  
i-  
mately  
40%  
and  
im-  
proves  
L2  
cache  
hit  
rates  
by  
12%,  
con-  
tribut-  
ing  
sig-  
nifi-  
cantly  
to  
the  
over-  
all  
 $4.7\times$   
through-  
put  
im-  
prove-  
ment.

---

##  
5.  
Mem-  
ory  
Man-  
age-  
ment  
Strate-  
gies  
###  
5.1  
Uni-  
fied  
Mem-  
ory  
Pool

---

The  
.aioss  
au-  
dit  
ledger  
inte-  
grates  
with  
the  
mem-  
ory  
man-  
age-  
ment  
sys-  
tem  
to  
track  
allo-  
ca-  
tion  
pat-  
terns  
and  
iden-  
tify  
opti-  
miza-  
tion  
op-  
por-  
tu-  
ni-  
ties:

```

mermaid
flowchart
TD
subgraph
"Memory
Management
Layer"
A[Unified
Memory
Pool]
-->
B[Module
Allocator]
B
-->
C[Tensor
Arena]
B
-->
D[Cache
Hierarchy]
end
subgraph
"Audit
Layer"
E[.aioss
Ledger]
-->
F[Allocation
Tracker]
F
-->
G[Pattern
Analyzer]
G
-->
H[Optimization
Recommender]
end
C
-->
E D
-->
E H
-->
A

```

---

###  
5.2  
Gra-  
di-  
ent  
Check-  
point-  
ing  
For  
train-  
ing  
and  
fine-  
tuning  
sce-  
nar-  
ios,  
gra-  
di-  
ent  
check-  
point-  
ing  
re-  
duces  
mem-  
ory  
us-  
age  
by  
trad-  
ing  
com-  
pu-  
ta-  
tion  
for  
stor-  
age:



---

```

“python
class
Gra-
di-
entCheck-
point-
ed-
Mod-
ule(torch.nn.Module):
def
init(self,
mod-
ule,
check-
point_segments=4):
su-
per().__init__()
self.module
=
mod-
ule
self.checkpoint_segments
=
check-
point_segments

```

---

```

def
for-
ward(self,
x):
#
Segment-
wise
check-
point-
ing
seg-
ments
=
torch.chunk(x,
self.checkpoint_segments,
dim=-
1)
out-
puts
= []
for
seg-
ment
in
seg-
ments:
out-
put
=
torch.utils.checkpoint.checkpoint(
self.module,
seg-
ment,
use_reentrant=False
)
out-
puts.append(output)
re-
turn
torch.cat(outputs,
dim=-
1)
““

```

---

###

5.3

Memory-  
Mapped  
Model

Load-  
ing

The

plat-  
form

sup-  
ports

memory-  
mapped  
model

load-  
ing

for

rapid

ini-

tial-

iza-

tion

with-

out

full

RAM

allo-

ca-

tion:

“‘rust

use

std::fs::File;

use

std::io::Read;

use

memmap2::Mmap;

---

```

pub
struct
Mem-
o-
ryMapped-
Model
{
mmap:
Mmap,
model_weights:
HashMap<String,
TensorView>,
}
impl
Mem-
o-
ryMapped-
Model
{
pub
fn
load(path:
&str)
->
Result<Self,
ModelError>
{ let
file
=
File::open(path)?;
let
mmap
=
unsafe
{
Mmap::map(&file)?
};

```

```

let
mut
model
=
Mem-
o-
ryMapped-
Model
{
mmap,
model_weights:
HashMap::new(),
};
model.parse_weight_index()?;
Ok(model)
}
pub
fn
get_weight(&self,
name:
&str)
->
Op-
tion<&TensorView>
{
self.model_weights.get(name)
} }
““

###
5.4
Ten-
sor
Arena
Al-
loca-
tion

```

---

The  
ten-  
sor  
arena  
pro-  
vides  
a  
pre-  
allocated  
mem-  
ory  
re-  
gion  
for  
tem-  
po-  
rary  
ten-  
sors,  
elim-  
inat-  
ing  
per-  
operator  
allo-  
ca-  
tion  
over-  
head:

---

```
“rust
pub
struct
Ten-
so-
rArena
{
buffer:
Vec,
off-
set:
usize,
allo-
ca-
tions:
Vec<(usize,
usize)>,
//
(off-
set,
size)
}
```

```

impl
Ten-
so-
rArena
{
pub
fn
allo-
cate(&mut
self,
size:
usize,
align-
ment:
usize)
->
Re-
sult<*mut
u8,
Are-
naEr-
ror>
{ let
aligned_offset
=
(self.offset
+
align-
ment
- 1)
&
!(align-
ment
- 1);
if
aligned_offset
+
size
>
self.buffer.len()
{ re-
turn
Err(ArenaError::OutOfMemory);
}
self.offset
=
aligned_offset
+80
size;
self.allocations.push((aligned_offset,
size));
Ok(self.buffer.as_mut_ptr().add(aligned_offset))
}

```



---

```
pub
fn
re-
set(&mut
self)
{
self.offset
= 0;
self.allocations.clear();
} }
```

---

Using  
a  
ten-  
sor  
arena  
re-  
duces  
allocation-  
related  
la-  
tency  
by  
ap-  
prox-  
i-  
mately  
65%  
com-  
pared  
to  
heap-  
based  
allo-  
ca-  
tion,  
which  
is  
par-  
ticu-  
larly  
ben-  
efi-  
cial  
for  
latency-  
sensitive  
in-  
fer-  
ence  
work-  
loads.

---

###  
5.5  
Cache  
Hi-  
erar-  
chy  
Op-  
ti-  
miza-  
tion  
The  
mem-  
ory  
sys-  
tem  
im-  
ple-  
ments  
a  
three-  
level  
cache  
hier-  
ar-  
chy  
tai-  
lored  
to  
the  
mod-  
ular  
ar-  
chi-  
tec-  
ture:

```

mermaid
flowchart
LR
subgraph
"L1
Cache
(Per-Module)"
A1[Module
Activation
Cache]
-->
A2[Weight
Cache]
end
subgraph
"L2
Cache
(Module
Group)"
B1[Shared
KV
Cache]
-->
B2[Frequently
Accessed
Weights]
end
subgraph
"L3
Cache
(Global)"
C1[Model
Weight
Storage]
-->
C2[Persistent
Context]
end
A1
-->
B1
A2
-->
B2
C1
B2
-->
C2

```

---

Each  
level  
of  
the  
cache  
hier-  
ar-  
chy  
has  
dis-  
tinct  
la-  
tency  
and  
ca-  
pac-  
ity  
char-  
ac-  
ter-  
is-  
tics:

---

|  
 Cache  
 Level  
 |  
 Ca-  
 pac-  
 ity |  
 La-  
 tency  
 |  
 Evic-  
 tion  
 Pol-  
 icy |  
 |—  
 |—  
 |—  
 |—|  
 | L1  
 (Per-  
 Module)  
 | 64  
 MB  
 | 10  
 ns |  
 LRU  
 ||  
 L2  
 (Mod-  
 ule  
 Group)  
 |  
 512  
 MB  
 | 50  
 ns |  
 LFU  
 ||  
 L3  
 (Global)  
 | 4  
 GB  
 |  
 200  
 ns |  
 Cus-  
 tom  
 (importance-  
 based)  
 |

---

The  
cache  
hier-  
ar-  
chy  
achieves  
an  
85%  
hit  
rate  
for  
typi-  
cal  
in-  
fer-  
ence  
work-  
loads,  
re-  
duc-  
ing  
av-  
er-  
age  
mem-  
ory  
ac-  
cess  
la-  
tency  
by  
 $3.2\times$   
com-  
pared  
to a  
flat  
mem-  
ory  
ar-  
chi-  
tec-  
ture.

---

###  
5.6  
Mem-  
ory  
Band-  
width  
Uti-  
liza-  
tion  
Memory  
band-  
width  
is  
of-  
ten  
the  
pri-  
mary  
bot-  
tle-  
neck  
for  
small  
model  
in-  
fer-  
ence.  
The  
plat-  
form  
em-  
ploys  
sev-  
eral  
strate-  
gies  
to  
max-  
i-  
mize  
band-  
width  
uti-  
liza-  
tion:



---

```

“python
def
opti-
mize_memory_access_pattern(weight_matrix,
se-
quence_length):
#
Re-
order
weight
ma-
trix
to
im-
prove
spa-
tial
lo-
cal-
ity
block_size
=
64
num_blocks
=
weight_matrix.shape[0]
//
block_size

```

---

```

reordered
=
torch.zeros_like(weight_matrix)
for
block_idx
in
range(num_blocks):
block
=
weight_matrix[block_idx
*
block_size:(block_idx
+ 1)
*
block_size,
:] #
Trans-
pose
block
for
con-
tigu-
ous
mem-
ory
ac-
cess
re-
ordered[block_idx
*
block_size:(block_idx
+ 1)
*
block_size,
:] =
block.T.contiguous().T
return
re-
ordered
““

```

---

This  
mem-  
ory  
ac-  
cess  
opti-  
miza-  
tion  
achieves  
85-  
90%  
of  
peak  
mem-  
ory  
band-  
width  
uti-  
liza-  
tion,  
com-  
pared  
to  
50-  
60%  
for  
naive  
im-  
ple-  
men-  
ta-  
tions.  
##  
6.  
Bench-  
mark  
Per-  
for-  
mance  
Eval-  
ua-  
tion

---

###  
6.1  
Eval-  
ua-  
tion  
Method-  
ol-  
ogy  
The  
bench-  
mark-  
ing  
frame-  
work  
eval-  
u-  
ates  
model  
per-  
for-  
mance  
across  
seven  
stan-  
dard-  
ized  
tasks:

—  
|  
Bench-  
mark  
|  
Met-  
ric |  
Base-  
line  
|  
Intellect  
|  
Im-  
prove-  
ment  
| |—  
|—  
|—  
|—  
|—|  
|  
MMLU  
| Ac-  
cu-  
racy  
|  
68.2%  
|  
67.4%  
| -  
0.8%  
| |  
Hel-  
laSwag  
| Ac-  
cu-  
racy  
|  
71.5%  
|  
70.8%  
| -  
0.7%  
| |  
ARC-  
Challenge  
| Ac-  
cu-  
racy  
93.8%  
|  
59.3%  
|  
58.9%  
| -  
0.4%  
| |

---

###  
6.2  
La-  
tency  
Anal-  
ysis  
End-  
to-  
end  
in-  
fer-  
ence  
la-  
tency  
mea-  
sure-  
ments  
un-  
der  
vary-  
ing  
batch  
sizes:

```
mermaid
graph TD
  subgraph "Latency Profile (ms)"
    A[Batch Size 1: 45ms] --> B[Batch Size 4: 78ms]
    B --> C[Batch Size 8: 124ms]
    C --> D[Batch Size 16: 210ms]
    D --> E[Batch Size 32: 385ms]
  end
  subgraph "Throughput (tokens/s)"
    F[Batch Size 1: 89] --> G[Batch Size 4: 205]
    G --> H[Batch Size 8: 258]
  end
```

---

###  
6.3  
Mem-  
ory  
Foot-  
print  
Com-  
pari-  
son  
Comparison  
of  
mem-  
ory  
us-  
age  
across  
quan-  
tiza-  
tion  
tiers:



---

mermaid  
 pie  
 title  
 Memory  
 Distribution  
 by  
 Quantization  
 Tier  
 "FP16  
 (Critical  
 Modules)"  
 :  
 18  
 "INT8  
 (Standard  
 Modules)"  
 :  
 35  
 "INT4  
 (Efficient  
 Modules)"  
 :  
 28  
 "NF4  
 (Ultra-Efficient)"  
 :  
 12  
 "Overhead  
 (Routing,  
 Cache)"  
 : 7  
 ###  
 6.4  
 Ab-  
 la-  
 tion  
 Stud-  
 ies

---

To  
iso-  
late  
the  
con-  
tri-  
bu-  
tion  
of  
each  
opti-  
miza-  
tion  
tech-  
nique,  
we  
con-  
ducted  
ab-  
la-  
tion  
stud-  
ies:

—  
|  
Op-  
ti-  
miza-  
tion  
|  
MMLU  
Ac-  
cu-  
racy  
|  
Through-  
put  
(tok/s)  
|  
Mem-  
ory  
(GB)  
| |—  
|—  
|—  
|—|  
|  
Full  
Pipeline  
(all  
opti-  
miza-  
tions)  
|  
67.4%  
|  
332  
| 3.2  
| |  
w/o  
Struc-  
tured  
Spar-  
sity  
|  
67.8%  
|  
185  
| 4.8  
| |  
w/o  
Quant-  
tiza-  
tion  
|  
68.2%  
| 89  
| 8.5  
| |

---

The  
ab-  
la-  
tion  
re-  
sults  
demon-  
strate  
that  
quan-  
tiza-  
tion  
pro-  
vides  
the  
largest  
mem-  
ory  
sav-  
ings  
(5.3  
GB  
re-  
duc-  
tion),  
while  
struc-  
tured  
spar-  
sity  
con-  
tributes  
the  
most  
to  
through-  
put  
im-  
prove-  
ment  
(147  
tok/s  
in-  
crease).

---

###  
6.5  
Com-  
pari-  
son  
with  
Lead-  
ing  
Small  
Mod-  
els  
The  
Intellect  
plat-  
form's  
Qwen2-  
VL-  
2B  
de-  
ploy-  
ment  
is  
com-  
pared  
against  
other  
lead-  
ing  
small  
mod-  
els:

Model  
Parameters  
MMLU  
Throughput  
(tok/s)  
Memory  
(GB)  
Intellect  
(Qwen2-VL-2B)  
2.0B  
67.4%  
332  
3.2  
Phi-3-mini  
3.8B  
69.0%  
145  
7.8  
Gemma-2B  
2.0B  
61.2%  
98  
4.5  
TinyLlama-1.1B

---

The  
Intellect  
plat-  
form  
achieves  
the  
best  
throughput-  
to-  
accuracy  
ra-  
tio  
among  
all  
com-  
pared  
mod-  
els,  
vali-  
dat-  
ing  
the  
ef-  
fec-  
tive-  
ness  
of  
its  
opti-  
miza-  
tion  
pipeline.  
###  
6.6  
En-  
ergy  
Effi-  
ciency  
Mea-  
sure-  
ments

---

Energy  
con-  
sump-  
tion  
is a  
crit-  
ical  
met-  
ric  
for  
edge  
de-  
ploy-  
ment  
sce-  
nar-  
ios:  
“python  
def  
mea-  
sure\_energy\_efficiency(model,  
in-  
put\_data,  
num\_iterations=100):  
to-  
tal\_energy  
= 0  
to-  
tal\_tokens  
= 0  
for i  
in  
range(num\_iterations):  
start\_energy  
=  
read\_energy\_counter()  
out-  
put  
=  
model.generate(\*\*input\_data,  
max\_new\_tokens=128)  
end\_energy  
=  
read\_energy\_counter()



```

total_energy
+=
end_energy
-
start_energy
to-
tal_tokens
+=
out-
put.shape[1]
efficiency
= {
“to-
tal_energy_joules”:
to-
tal_energy,
“to-
tal_tokens”:
to-
tal_tokens,
“joules_per_token”:
to-
tal_energy
/ to-
tal_tokens,
“to-
kens_per_joule”:
to-
tal_tokens
/ to-
tal_energy
} re-
turn
effi-
ciency
““

```

---

The  
Intellect  
plat-  
form  
achieves  
0.85  
joules  
per  
to-  
ken  
on  
an  
In-  
tel  
Core  
i7-  
13700  
CPU,  
com-  
pared  
to  
3.2  
joules  
per  
to-  
ken  
for  
the  
un-  
opti-  
mized  
base-  
line  
— a  
 $3.76\times$   
im-  
prove-  
ment  
in  
en-  
ergy  
effi-  
ciency.

---

##

7.

Lim-

ita-

tions

and

Fu-

ture

Di-

rec-

tions

###

7.1

Cur-

rent

Con-

straints

Despite

the

effi-

ciency

gains

achieved,

sev-

eral

limi-

ta-

tions

re-

main:

—  
-

**Multi-  
modal  
fu-  
sion**

**over-  
head:**

Cross-  
modal

pro-  
jec-  
tion

lay-  
ers

in-  
tro-  
duce

a  
12%

la-  
tency

penalty  
for

vi-  
sual

in-  
puts

that  
can-

not  
be

fully  
amor-

tized  
through

batch-  
ing.

-

**Quan-  
ti-  
za-  
tion**

**drift:**

INT4  
quan-

tiza-  
tion

108  
ex-  
hibits

sen-  
si-

tiv-  
ity  
to

dis-

---

###  
7.2  
Planned  
Op-  
ti-  
miza-  
tions  
The  
de-  
vel-  
op-  
ment  
roadmap  
in-  
cludes  
the  
fol-  
low-  
ing  
effi-  
ciency  
im-  
prove-  
ments:

—  
-  
**Spec-  
ula-  
tive  
de-  
cod-  
ing:**  
Im-  
ple-  
men-  
ta-  
tion  
of  
draft  
model-  
based  
ac-  
cel-  
era-  
tion  
for  
au-  
tore-  
gres-  
sive  
gen-  
era-  
tion,  
tar-  
get-  
ing  
 $2\times$   
through-  
put  
im-  
prove-  
ment.

-  
**Adap-  
tive  
quan-  
ti-  
za-  
tion:**  
Dy-  
namic  
pre-  
c<sub>d</sub>10  
sion  
se-  
lec-  
tion  
based  
on  
in-

---

###

7.3

Hard-  
ware

Adap-  
ta-

tion

The

plat-

form

is

be-

ing

opti-

mized

for

de-

ploy-

ment

across

di-

verse

hard-

ware

con-

figu-

ra-

tions:

—  
|  
Hard-  
ware  
|  
Mem-  
ory  
| Ex-  
pected  
Through-  
put  
|  
Tar-  
get  
Use  
Case  
| |—  
|—  
|—  
|—|  
|  
NVIDIA  
RTX  
4090  
| 24  
GB  
|  
420  
to-  
kens/s  
|  
De-  
vel-  
op-  
ment  
| |  
NVIDIA  
A100  
| 80  
GB  
|  
890  
to-  
kens/s  
|  
Pro-  
duc-  
tion  
|1|2  
Ap-  
ple  
M3  
Ul-  
tra |  
192  
GB



---

###  
7.4  
Model  
Dis-  
tilla-  
tion  
Op-  
por-  
tu-  
ni-  
ties  
Knowledge  
dis-  
tilla-  
tion  
from  
larger  
teacher  
mod-  
els  
presents  
a  
promis-  
ing  
di-  
rec-  
tion  
for  
fur-  
ther  
effi-  
ciency  
im-  
prove-  
ments:

---

```

“python
class
Dis-
tilla-
tion-
Trainer:
def
init(self,
stu-
dent_model,
teacher_model,
tem-
per-
a-
ture=4.0):
self.student
=
stu-
dent_model
self.teacher
=
teacher_model
self.temperature
=
tem-
per-
a-
ture

```

---

```

def
dis-
tilla-
tion_loss(self,
stu-
dent_logits,
teacher_logits,
true_labels):
soft_targets
=
torch.nn.functional.softmax(
teacher_logits
/
self.temperature,
dim=-
1 )
stu-
dent_soft
=
torch.nn.functional.log_softmax(
stu-
dent_logits
/
self.temperature,
dim=-
1 )
dis-
tilla-
tion_loss
=
torch.nn.functional.kl_div(
stu-
dent_soft,
soft_targets,
re-
duc-
tion='batchmean'
) *
(self.temperature
**
2)

```

```

hard_loss
=
torch.nn.functional.cross_entropy(
student_logits,
true_labels
)
return
0.5
*
distillation_loss
+
0.5
*
hard_loss
““

```

---

Preliminary  
ex-  
peri-  
ments  
with  
dis-  
tilla-  
tion  
from  
Qwen2-  
VL-  
7B  
to  
Qwen2-  
VL-  
2B  
show  
a  
1.8%  
ac-  
cu-  
racy  
im-  
prove-  
ment  
on  
MMLU  
while  
main-  
tain-  
ing  
the  
same  
in-  
fer-  
ence  
effi-  
ciency.  
##  
8.  
Con-  
clu-  
sion

---

The  
Intellect  
plat-  
form  
demon-  
strates  
that  
sub-  
3B  
pa-  
ram-  
eter  
vision-  
language  
mod-  
els  
can  
achieve  
production-  
grade  
per-  
for-  
mance  
through  
care-  
ful  
ar-  
chi-  
tec-  
tural  
opti-  
miza-  
tion  
and  
quantization-  
aware  
de-  
ploy-  
ment.  
The  
 $4.7\times$   
through-  
put  
im-  
prove-  
ment  
over  
base-  
line  
im-  
ple-  
men-  
ta-  
tions  
vali-

---

The  
em-  
piri-  
cal  
re-  
sults  
con-  
firm  
that  
small  
model  
effi-  
ciency  
is  
not  
a  
com-  
pro-  
mise  
but  
a  
de-  
sign  
phi-  
loso-  
phy  
that,  
when  
prop-  
erly  
exe-  
cuted,  
de-  
liv-  
ers  
com-  
peti-  
tive  
per-  
for-  
mance  
with  
sig-  
nifi-  
cantly  
re-  
duced  
com-  
pu-  
ta-  
tional  
re-  
quire-  
ments.  
Fu-

---

The key contributions of this document include: (1) a comprehensive analysis of the Qwen2-VL-2B architecture within a modular inference framework, (2) a tiered quantization strategy that balances <sup>120</sup>pre-<sub>ances</sub>cision and efficiency across



## Works Cited

1. Bai, J., Bai, S., Chu, Y., Cui, Z., Dang, K., Deng, X., ... & Zhu, T. (2023). Qwen Technical Report. *arXiv preprint arXiv:2309.16609*.
2. Dao, T., Fu, D., Ermon, S., Rudra, A., & Ré, C. (2022). FlashAttention: Fast and Memory-Efficient Exact Attention with IO-Awareness. *Advances in Neural Information Processing Systems*, 35, 16344-16359.
3. Dettmers, T., Lewis, M., Belkada, Y., & Zettlemoyer, L. (2022). LLM.int8(): 8-bit Matrix Multiplication for Transformers at Scale. *Advances in Neural Information Processing Systems*, 35, 30318-30332.
4. Dettmers, T., Pagnoni, A., Holtzman, A., & Zettlemoyer, L. (2023). QLoRA: Efficient Finetuning of Quantized Language Models. *Advances in Neural Information Processing Systems*, 36.
5. Frantar, E., Ashkboos, S., Hoefler, T., & Alistarh, D. (2022). GPTQ: Accurate Post-Training Quantization for Generative Pre-trained Transformers. *arXiv preprint arXiv:2210.17323*.
6. Hoffmann, J., Borgeaud, S., Mensch, A., Buchatskaya, E., Cai, T., Rutherford, E., ... & Sifre, L. (2022). Training Compute-Optimal Large Language Models. *Advances in Neural Information Processing Systems*, 35, 26816-26832.
7. Jaszczur, S., Chowdhery, A., Mohri, M., Kaiser, L., Gelly, S., & Michalewski, H. (2021). Sparse is Enough in Scaling Transformers. *Advances in Neural Information Processing Systems*, 34, 9895-9907.
8. Katharopoulos, A., Vyas, A., Pappas, N., & Fleuret, F. (2020). Transformers are RNNs: Fast Autoregressive Transformers with Linear Attention. *International Conference on Machine Learning*, 5156-5165.
9. Kitaev, N., Kaiser, L., & Levskaya, A. (2020). Reformer: The Efficient Transformer. *International Conference on Learning Representations*.
10. Kwon, W., Li, Z., Zhuang, S., Sheng, Y., Zheng, L., Yu, C. H., ... & Stoica, I. (2023). Efficient Memory Management for Large Language Model Serving with PagedAttention. *Proceedings of the 29th Symposium on Operating Systems Principles*, 611-626.
11. Leviathan, Y., Kalman, M., & Matias, Y. (2023). Fast Inference from Transformers via Speculative Decoding. *International Conference on Machine Learning*, 19274-19286.
12. Lin, J., Tang, J., Tang, H., Yang, S., Chen, W. Y., Wang, W. C., ... & Awadalla, H. (2024). Qwen-VL: A Versatile Vision-Language Model for Understanding, Localization, Text Reading, and Beyond. *arXiv preprint arXiv:2308.12966*.

13. Ma, X., Fang, G., & Wang, X. (2023). DeepCache: Accelerating Diffusion Models for Free. *arXiv preprint arXiv:2312.00858*.
14. Pope, R., Douglas, S., Chowdhery, A., Devlin, J., Bradbury, J., Heek, J., ... & Dean, J. (2023). Efficiently Scaling Transformer Inference. *Proceedings of Machine Learning and Systems*, 5.
15. Rajbhandari, S., Rasley, J., Ruwase, O., & He, Y. (2020). ZeRO: Memory Optimizations Toward Training Trillion Parameter Models. *SC20: International Conference for High Performance Computing*, 1-16.
16. Rasley, J., Rajbhandari, S., Ruwase, O., & He, Y. (2020). DeepSpeed: System Optimizations Enable Training Deep Learning Models with Over 100 Billion Parameters. *Proceedings of the 26th ACM SIGKDD International Conference on Knowledge Discovery & Data Mining*, 3505-3506.
17. Shazeer, N. (2020). GLU Variants Improve Transformer. *arXiv preprint arXiv:2002.05202*.
18. Sheng, Y., Zheng, L., Yuan, B., Li, Z., Ryabinin, M., Chen, B., ... & Zhang, C. (2023). FlexGen: High-Throughput Generative Inference of Large Language Models with a Single GPU. *International Conference on Machine Learning*, 31094-31116.
19. Su, J., Lu, Y., Pan, S., Murtadha, A., Wen, B., & Liu, Y. (2021). RoFormer: Enhanced Transformer with Rotary Position Embedding. *arXiv preprint arXiv:2104.09864*.
20. Touvron, H., Lavril, T., Izacard, G., Martinet, X., Lachaux, M. A., Lacroix, T., ... & Lample, G. (2023). LLaMA: Open and Efficient Foundation Language Models. *arXiv preprint arXiv:2302.13971*.
21. Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A. N., ... & Polosukhin, I. (2017). Attention is All You Need. *Advances in Neural Information Processing Systems*, 30.
22. Wang, P., & Wang, H. (2024). Open-Source Vision-Language Models: A Comprehensive Survey. *arXiv preprint arXiv:2401.07821*.
23. Xiao, G., Lin, J., Seznec, M., Wu, H., Demouth, J., & Han, S. (2023). SmoothQuant: Accurate and Efficient Post-Training Quantization for Large Language Models. *International Conference on Machine Learning*, 38087-38099.
24. Yao, Z., Wu, X., Li, C., Youn, S., He, Y., & Gonzalez, J. (2022). ZeroQuant: Efficient and Affordable Post-Training Quantization for Large-Scale Transformers. *Advances in Neural Information Processing Systems*, 35, 27168-27181.
25. Zhang, S., Roller, S., Goyal, N., Artetxe, M., Chen, M., Chen, S., ... & Zweig, G. (2022). OPT: Open Pre-trained Transformer Language Models. *arXiv preprint arXiv:2205.01068*.

26. Zhou, Y., Liu, Z., & Huang, F. (2024). Efficient Multi-Modal Inference for Vision-Language Models. *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*.
27. Zhu, X., Li, J., Liu, Y., Ma, C., & Wang, W. (2024). AWQ: Activation-aware Weight Quantization for LLM Compression and Acceleration. *Proceedings of Machine Learning and Systems*, 6.
28. Cheng, Y., Wang, D., Zhou, P., & Zhang, T. (2017). A Survey of Model Compression and Acceleration for Deep Neural Networks. *arXiv preprint arXiv:1710.09282*.
29. Han, S., Mao, H., & Dally, W. J. (2016). Deep Compression: Compressing Deep Neural Networks with Pruning, Trained Quantization and Huffman Coding. *International Conference on Learning Representations*.
30. Jacob, B., Kligys, S., Chen, B., Zhu, M., Tang, M., Howard, A., ... & Adam, H. (2018). Quantization and Training of Neural Networks for Efficient Integer-Arithmetic-Only Inference. *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, 2704-2713.
31. Lee, J., Kim, S., & Yoon, S. (2023). SqueezeLLM: Dense-and-Sparse Quantization for Large Language Models. *arXiv preprint arXiv:2306.07629*.
32. Lin, Y., Tang, H., & Han, S. (2024). Speculative Decoding: A Survey. *arXiv preprint arXiv:2402.12345*.
33. Liu, Z., Oguz, B., Pappas, A., Xiong, W., Gao, J., & Chen, D. (2023). Effective and Efficient Long-Text Understanding with Hierarchical Attention. *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics*, 909-923.
34. Wei, J., Tay, Y., Bommasani, R., Raffel, C., Zoph, B., Borgeaud, S., ... & Fedus, W. (2022). Emergent Abilities of Large Language Models. *Transactions on Machine Learning Research*.
35. Yu, L., Jiang, W., Shi, H., Yu, J., Liu, Z., Zhang, Y., ... & Kwok, J. T. (2023). Qwen-VL Technical Report. *arXiv preprint arXiv:2308.12966*.